
When are Non-Parametric Methods Robust?

Robi Bhattacharjee^{* 1} Kamalika Chaudhuri^{* 1}

Abstract

A growing body of research has shown that many classifiers are susceptible to *adversarial examples* – small strategic modifications to test inputs that lead to misclassification. In this work, we study general non-parametric methods, with a view towards understanding when they are robust to these modifications. We establish general conditions under which non-parametric methods are r -consistent – in the sense that they converge to optimally robust and accurate classifiers in the large sample limit.

Concretely, our results show that when data is well-separated, nearest neighbors and kernel classifiers are r -consistent, while histograms are not. For general data distributions, we prove that pre-processing by Adversarial Pruning (Yang et al., 2019) – that makes data well-separated – followed by nearest neighbors or kernel classifiers also leads to r -consistency.

1. Introduction

Recent work has shown that many classifiers tend to be highly non-robust and that small strategic modifications to regular test inputs can cause them to misclassify (Szegedy et al., 2014; Goodfellow et al., 2014; Lowd & Meek, 2005). Motivated by the use of machine learning in safety-critical applications, this phenomenon has recently received considerable interest; however, what exactly causes this phenomenon – known in the literature as *adversarial examples* – still remains a mystery.

Prior work has looked at three plausible reasons why adversarial examples might exist. The first, of course, is the possibility that in real data distributions, different classes are very close together in space – which does not seem plausible in practice. Another possibility is that classification algorithms may require more data to be robust than to be merely accurate; some prior work (Schmidt et al., 2018; Wang et al.,

2018; Montasser et al., 2019) suggests that this might be true for certain classifiers or algorithms. Finally, others (Salman et al., 2019; Avdiukhin et al., 2019; Wang et al., 2018) have suggested that better training algorithms may give rise to more robust classifiers – and that in some cases, finding robust classifiers may even be computationally challenging.

In this work, we consider this problem in the context of general non-parametric classifiers. Contrary to parametrics, non-parametric methods are a form of local classifiers, and include a large number of pattern recognition methods such as nearest neighbors, decision trees, random forests and kernel classifiers. There is a richly developed statistical theory of non-parametric methods (Devroye et al., 1996), which focuses on accuracy, and provides very general conditions under which these methods converge to the Bayes optimal with growing number of samples. We, in contrast, analyze robustness properties of these methods, and ask instead when they converge to the classifier with the highest astuteness at a desired radius r . Recall that the astuteness of a classifier at radius r is the fraction of points from the distribution on which it is accurate and has the same prediction up to a distance r (Wang et al., 2018; Schmidt et al., 2018).

We begin by looking at the very simple case when data from different classes is well-separated – by at least a distance $2r$. Although achieving astuteness in this case may appear trivial, we show that even in this highly favorable case, not all non-parametric methods provide robust classifiers – and this even holds for methods that converge to the Bayes optimal in the large sample limit.

This raises the natural question – when do non-parametric methods produce astute classifiers? We next provide conditions under which a non-parametric method converges to the most astute classifier in the large sample limit under well-separated data. Our conditions are analogous to the classical conditions for convergence to the Bayes optimal (Devroye et al., 1996; Stone, 1977), but a little stronger. We show that nearest neighbors and kernel classifiers whose kernel functions decay fast enough, satisfy these conditions, and hence converge to astute classifiers in the large sample limit. In contrast, histogram classifiers, which do converge to the Bayes optimal in the large sample limit, may not converge to the most astute classifier. This indicates that there may be some non-parametric methods, such as nearest neighbors and kernel classifiers, that are more naturally robust when

^{*}Equal contribution ¹Department of Computer Science, University of California, San Diego. Correspondence to: <>.

trained on well-separated data, and some that are not.

What happens when different classes in the data are not as well-separated? For this case, (Yang et al., 2019) proposes a method called Adversarial Pruning that preprocesses the training data by retaining the maximal set of points such that different classes are distance $\geq 2r$ apart, and then trains a non-parametric method on the pruned data. We next prove that if a non-parametric method has certain properties, then the classifier produced by Adversarial Pruning followed by the method does converge to the most astute classifier in the large sample limit. We show that again nearest neighbors and kernel classifiers whose kernel functions decay faster than inverse polynomials satisfy these properties. Our results thus complement and build upon the empirical results of (Yang et al., 2019) by providing a performance guarantee.

What can we conclude about the cause for adversarial examples? Our results seem to indicate that at least for non-parametrics, it is mostly the training algorithms that are responsible. With a few exceptions, decades of prior work in machine learning and pattern recognition has largely focussed on designing training methods that provide increasingly accurate classifiers – perhaps to the detriment of other aspects such as robustness. In this context, our results serve to (a) provide a set of guidelines that can be used for designing non-parametric methods that are robust and accurate on well-separated data and (b) demonstrate that when data is not well-separated, preprocessing through adversarial pruning (Yang et al., 2019) may be used to ensure convergence to optimally astute solutions in the large sample limit.

1.1. Related Work

There is a large body of work on adversarial attacks (Carlini & Wagner, 2017; Liu et al., 2017; Papernot et al., 2017; 2016a; Szegedy et al., 2014) and defenses (Hein & Andriushchenko, 2017; Katz et al., 2017; Madry et al., 2018; Papernot et al., 2016b; Raghuathan et al., 2018; Sinha et al., 2018) in the parametric setting, specifically focusing on neural networks. On the other hand, adversarial examples for nonparametric classifiers have mostly been studied in a much more ad-hoc manner, and to our knowledge, there has been no theoretical investigation into general properties of algorithms that promote robustness in non-parametric classifiers.

For nearest neighbors, there has been some prior work on adversarial attacks (Amsaleg et al., 2017; Sitawarin & Wagner, 2019; Wang et al., 2018; Yang et al., 2019) as well as defenses. Wang et al. (Wang et al., 2018) proposes a defense for 1-NN by pruning the input sample. However, their defense learns a classifier whose robustness regions converge towards those of the Bayes optimal classifier, which itself may potentially have poor robustness properties. Yang et al. (Yang et al., 2019) accounts for this problem by proposing

the notion of the r -optimal classifier, and propose an algorithm called Adversarial Pruning which can be interpreted as a finite sample approximation to the r -optimal. However, they do not provide formal performance guarantees for Adversarial Pruning, which we do.

For Kernel methods, Hein and Andriushchenko (Hein & Andriushchenko, 2017) study lower bounds on the norm of the adversarial manipulation that is required for changing a classifier’s output. They specifically study bounds for Kernel Classifiers, and propose an empirically based regularization idea that improves robustness. In this work, we improve the robustness properties of kernel classification through adversarial pruning, and show formal guarantees regarding convergence towards the r -optimal classifier.

For decision trees and random forests, attacks and defenses have been provided by (Andriushchenko & Hein, 2019; Kantchelian et al., 2015; Chen et al., 2019). Again, most of the work here is empirical in nature, and convergence guarantees are not provided.

Pruning has a long history of being applied for improving nearest neighbors (Gates, 1972; Gottlieb et al., 2014; Hart, 1968; Kontorovich et al., 2017; Kontorovich & Weiss, 2015; Hanneke et al., 2019), but this has been entirely done in the context of generalization, without accounting for robustness. In their work, Yang et al. empirically show that adversarial pruning can improve robustness for nearest neighbor classifiers. However, they do not provide any formal guarantees for their algorithms. In this work, we prove formal guarantees for *adversarial pruning* in the large sample limit, both for nearest neighbors as well as for more general *weight functions*.

There is a long history of literature for understanding the consistency of Kernel classifiers (Steinwart, 2005; Stone, 1977), but this has only been done for accuracy and generalization. In this work, we find different conditions are needed to ensure that a Kernel classifier converges in robustness in addition to accuracy.

2. Preliminaries

2.1. Setting

We consider binary classification where instances are drawn from a totally bounded metric space $\mathcal{X}$ that is equipped with distance metric denoted by d , and the label space is $\{\pm 1\} = \{-1, +1\}$. The classical goal of classification is to build a highly *accurate* classifier, which we define as follows.

Definition 1. (Accuracy) Let $\mathcal{D}$ be a distribution over $\mathcal{X} \times \{\pm 1\}$, and let $f \in \{\pm 1\}^{\mathcal{X}}$ be a classifier. Then the *accuracy* of f over $\mathcal{D}$, denoted $A(f, \mathcal{D})$, is the fraction of

examples $(x, y) \sim \mathcal{D}$ for which $f(x) = y$. Thus

$$A(f, \mathcal{D}) = P_{(x,y) \sim \mathcal{D}}[f(x) = y].$$

In this work, we consider *robustness* in addition to accuracy. Let $B(x, r)$ denoted the closed ball of radius r centered at x .

Definition 2. (*Robustness*) A classifier $f \in \{\pm 1\}^{\mathcal{X}}$ is said to be **robust** at x with radius r if $f(x) = f(x')$ for all $x' \in B(x, r)$.

Our goal is to find non-parametric algorithms that output classifiers that are robust, in addition to being accurate. To account for both criteria, we combine them into a notion of *astuteness* (Wang et al., 2018; Schmidt et al., 2018).

Definition 3. (*Astuteness*) A classifier $f \in \{\pm 1\}^{\mathcal{X}}$ is said to be **astute** at (x, y) with radius r if f is robust at x with radius r and $f(x) = y$. The **astuteness** of f over $\mathcal{D}$, denoted $A_r(f, \mathcal{D})$, is the fraction of examples $(x, y) \sim \mathcal{D}$ for which f is astute at (x, y) with radius r . Thus

$$A_r(f, \mathcal{D}) = P_{(x,y) \sim \mathcal{D}}[f(x') = y, \forall x' \in B(x, r)].$$

It is worth noting that $A_0(f, \mathcal{D}) = A(f, \mathcal{D})$, since astuteness with radius 0 is simply the accuracy. For this reason, we will use $A_0(f, \mathcal{D})$ to denote accuracy from this point forwards.

2.2. Notions of Consistency

Traditionally, a classification algorithm is said to be consistent if as the sample size grows to infinity, the accuracy of the classifier it learns converges towards the best possible accuracy on the underlying data distribution. We next introduce and formalize an alternative form of consistency, called r -consistency, that applies to robust classifiers.

We begin with a formal definition of the Bayes Optimal Classifier – the most accurate classifier on a distribution – and consistency.

Definition 4. (*Bayes Optimal Classifier*) The **Bayes Optimal Classifier** on a distribution $\mathcal{D}$, denoted by g^* , is defined as follows. Let $\eta(x) = p_{\mathcal{D}}(+1|x)$. Then

$$g^*(x) = \begin{cases} +1 & \eta(x) \geq 0.5 \\ -1 & \eta(x) < 0.5 \end{cases}$$

It can be shown that g^* achieves the highest accuracy over $\mathcal{D}$ over all classifiers.

Definition 5. (*Consistency*) Let M be a classification algorithm over $\mathcal{X} \times \{\pm 1\}$. M is said to be **consistent** if for any $\mathcal{D}$ over $\mathcal{X} \times \{\pm 1\}$, and any ϵ, δ over $(0, 1)$, there exists N such that for $n \geq N$, with probability $1 - \delta$ over $S \sim \mathcal{D}^n$, we have:

$$A(M(S), \mathcal{D}) \geq A(g^*, \mathcal{D}) - \epsilon,$$

where g^* is the bayes optimal classifier for $\mathcal{D}$.

How can we incorporate robustness in addition to accuracy in this notion? A plausible way, as used in (Wang et al., 2018), is that the classifier should converge towards being astute where the Bayes Optimal classifier is astute. However, the Bayes Optimal classifier is not necessarily the most astute classifier and may even have poor astuteness. To see this, consider the following example.

Example 1 Consider $\mathcal{D}$ over $\mathcal{X} = [0, 1]$ such that $\mathcal{D}_{\mathcal{X}}$ is the uniform distribution and

$$p(y = 1|x) = \frac{1}{2} + \sin \frac{4\pi x}{r}.$$

For any point x , there exists $x_1, x_2 \in ([x - r, x + r] \cap [0, 1])$ such that $p(y = 1|x_1) > \frac{1}{2}$ and $p(y = 1|x_2) < \frac{1}{2}$. $A_r(g^*, r) = 0$. However, the classifier that always predicts $f(x) = +1$ does better. It is robust everywhere, and since $P_{(x,y) \sim \mathcal{D}}[y = +1] = \frac{1}{2}$, it follows that $A_r(f, \mathcal{D}) = \frac{1}{2}$.

This motivates the notion of the r -optimal classifier, introduced by (Yang et al., 2019), which is the classifier with maximum astuteness.

Definition 6. (*r -optimal classifier*) The **r -optimal classifier** of a distribution G denoted by g_r^* is the classifier with maximum astuteness. Thus

$$g_r^* = \arg \max_{f \in \{\pm 1\}^{\mathcal{X}}} A_r(f, \mathcal{D}).$$

We let $A_r^*(\mathcal{D})$ denote $A_r(g_r^*, \mathcal{D})$.

Observe that g_r^* is not necessarily unique. To account for this, we use $A_r^*(\mathcal{D})$ in our definition for r -consistency.

Definition 7. (*r -consistent*) Let M be a classification algorithm over $\mathcal{X} \times \{\pm 1\}$. M is said to be **r -consistent** if for any $\mathcal{D}$, any $\epsilon, \delta \in (0, 1)$, and $0 < \gamma < r$, there exists N such that for $n \geq N$, with probability $1 - \delta$ over $S \sim \mathcal{D}^n$,

$$A_{r-\gamma}(M(S), \mathcal{D}) \geq A_r^*(\mathcal{D}) - \epsilon.$$

if the above conditions hold for a specific distribution $\mathcal{D}$, we say that M is r -consistent with respect to $\mathcal{D}$.

Observe that in addition to the usual ϵ and δ , there is an extra parameter γ which measures the gap in the robustness radius. We may need this parameter as when classes are exactly $2r$ apart, we may not be able to find the exact robust boundary with only finite samples.

Our analysis will be centered around understanding what kinds of algorithms M provide highly astute classifiers for a given radius r . We begin by first considering the special case of r -separated distributions.

Definition 8. (*r-separated distributions*) A distribution $\mathcal{D}$ is said to be **r-separated** if there exist subsets $T^+, T^- \subset \mathcal{X}$ such that

1. $\mathbb{P}_{(x,y) \sim \mathcal{D}}[x \in T^y] = 1$.
2. $\forall x_1 \in T^+, \forall x_2 \in T^-, d(x_1, x_2) > 2r$.

Observe that if $\mathcal{D}$ is *r-separated*, $A_r(g_r^*, \mathcal{D}) = 1$.

2.3. Non-parametric Classifiers

Many non-parametric algorithms classify points by averaging labels over a local neighborhood from their training data. A very general form of this idea is encapsulated in *weight functions* – which is the general form we will use.

Definition 9. (Devroye et al., 1996) A **weight function** W is a non-parametric classifier with the following properties.

1. Given input $S = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\} \sim \mathcal{D}^n$, W constructs functions $w_1^S, w_2^S, \dots, w_n^S : \mathcal{X} \rightarrow [0, 1]$ such that for all $x \in \mathcal{X}$, $\sum_{i=1}^n w_i^S(x) = 1$. The functions w_i^S are allowed to depend on $x_1, x_2, \dots, x_n$ but must be independent of $y_1, y_2, \dots, y_n$.

2. W has output W_S defined as

$$W_S(x) = \begin{cases} +1 & \sum_{i=1}^n w_i^S(x) y_i > 0 \\ -1 & \sum_{i=1}^n w_i^S(x) y_i \leq 0 \end{cases}$$

As a result, $w_i^S(x)$ can be thought of as the weight that (x_i, y_i) has in classifying x .

Weight functions encompass a fairly extensive set of common non-parametric classifiers, which is the motivation for considering them. We now define several common non-parametric algorithms that can be construed as weight functions.

Definition 10. A **histogram classifier**, H , is a non-parametric classification algorithm over $\mathbb{R}^d \times \{\pm 1\}$ that works as follows. For a distribution $\mathcal{D}$ over $\mathbb{R} \times \{\pm 1\}$, H takes $S = \{(x_i, y_i) : 1 \leq i \leq n\} \sim \mathcal{D}^n$ as input. Let k_i be a sequence with $\lim_{i \rightarrow \infty} k_i = \infty$ and $\lim_{i \rightarrow \infty} \frac{k_i}{i} = 0$. H constructs a set of hypercubes $C = \{c_1, c_2, \dots, c_m\}$ as follows:

1. Initially $C = \{c\}$, where $S \subset c$.
2. For $c \in C$, if c contains more than k_n points of S , then partition c into 2^d equally sized hypercubes, and insert them into C .
3. Repeat step 2 until all cubes in C have at most k_n points.

For $x \in \mathbb{R}$ let $c(x)$ denote the unique cell in C containing x . If $c(x)$ doesn't exist, then $H_S(x) = -1$ by default. Otherwise,

$$H_S(x) = \begin{cases} +1 & \sum_{x_i \in c(x)} y_i > 0 \\ -1 & \sum_{x_i \in c(x)} y_i \leq 0 \end{cases}.$$

Histogram classifiers are weight functions in which all x_i contained within the same cell as x are given the same weight $w_i^S(x)$ in predicting x , while all other x_i are given weight 0.

Definition 11. A **kernel classifier** is a weight function W over $\mathcal{X} \times \{\pm 1\}$ constructed from function $K : \mathbb{R}^+ \cup \{0\} \rightarrow \mathbb{R}^+$ and some sequence $\{h_n\} \subset \mathbb{R}^+$ in the following manner. Given $S = \{(x_i, y_i)\} \sim \mathcal{D}^n$, we have

$$w_i^S(x) = \frac{K(\frac{d(x, x_i)}{h_n})}{\sum_{j=1}^n K(\frac{d(x, x_j)}{h_n})}.$$

Then, as above, W has output

$$W_S(x) = \begin{cases} +1 & \sum_{i=1}^n w_i^S(x) y_i > 0 \\ -1 & \sum_{i=1}^n w_i^S(x) y_i \leq 0 \end{cases}$$

Finally, we note that k_n -nearest neighbors is also a weight function; $w_i^S(x) = \frac{1}{k_n}$ if x_i is one of the k_n closest neighbors of x and 0 otherwise.

3. Warm Up: r-separated distributions

We begin by considering the case when the data distribution is *r-separated*; the more general case is considered in Section 4. While classifying *r-separated* distributions robustly may appear almost trivial, learning an arbitrary classifier does not necessarily produce an astute result. To see this, consider the following example of a histogram classifier – which is known to be consistent.

We let H denote the histogram classifier over $\mathbb{R}$.

Example 2 Consider the data distribution $\mathcal{D} = \mathcal{D}^+ \cup \mathcal{D}^-$ where $\mathcal{D}^+$ is the uniform distribution over $[0, \frac{1}{4})$ and $\mathcal{D}^-$ is the uniform distribution over $(\frac{1}{2}, 1]$, $p(+1|x) = 1$ for $x \in \mathcal{D}^+$, and $p(-1|x) = 1$ for $x \in \mathcal{D}^-$.

We make the following observations (refer to Figure 1).

1. $\mathcal{D}$ is 0.1-separated, since the supports of $\mathcal{D}^+$ and $\mathcal{D}^-$ have distance $0.25 > 0$.
2. If n is sufficiently large, H will construct the cell $[0.25, 0.5)$, which will not be split because it will never contain any points.
3. $H_S(x) = -1$ for $x \in [0.25, 0.5)$.

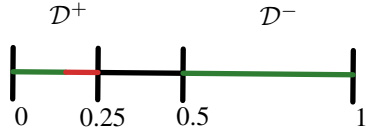


Figure 1. H_S is astute in the green region, but not robust in the red region.

4. H_S is not astute at $(x, 1)$ for $x \in (0.15, 0.25)$. Thus $A_{0.1}(H_S, \mathcal{D}) = 0.8$.

Example 2 shows that histogram classifiers do not always learn astute classifiers even when run on r -separated distributions. This motivates the question: which non-parametric classifiers do?

We answer this question in the following theorem, which gives sufficient conditions for a weight function (definition 9) to be r -consistent over an r -separated distribution.

Theorem 12. *Let $\mathcal{D}$ be a distribution over $\mathcal{X} \times \{\pm 1\}$, and let W be a weight function. Let X be a random variable with distribution $\mathcal{D}_X$, and $S = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\} \sim \mathcal{D}^n$. Suppose that for any $0 < a < b$,*

$$\lim_{n \rightarrow \infty} \mathbb{E}_{X, S} \left[\sup_{x' \in B(X, a)} \sum_{i=1}^n w_i^S(x') I_{||x_i - x'|| > b} \right] = 0.$$

Then if $\mathcal{D}$ is r -separated, W is r -consistent with respect to $\mathcal{D}$.

First, we compare Theorem 12 to Stone's theorem (Stone, 1977), which gives sufficient conditions for a weight function to be consistent (i.e. converge in accuracy towards the Bayes optimal). For convenience, we include a statement of Stone's theorem.

Theorem 13. (Stone, 1977) *Let W be weight function over $\mathcal{X} \times \{\pm 1\}$. Suppose the following conditions hold for any distribution $\mathcal{D}$ over $\mathcal{X} \times \{\pm 1\}$. Let X be a random variable with distribution $\mathcal{D}_X$, and $S = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\} \sim \mathcal{D}^n$. All expectations are taken over X and S .*

1. *There is a constant c such that, for every nonnegative measurable function f satisfying $\mathbb{E}[f(X)] < \infty$,*

$$\mathbb{E} \left[\sum_{i=1}^n w_i^S(X) f(x_i) \right] \leq c \mathbb{E}[f(X)].$$

2. *For all $a > 0$,*

$$\lim_{n \rightarrow \infty} \mathbb{E} \left[\sum_{i=1}^n w_i^S(x) I_{||x_i - X|| > a} \right] = 0,$$

where $I_{||x_i - X|| > a}$ is an indicator variable.

- 3.

$$\lim_{n \rightarrow \infty} \mathbb{E} \left[\max_{1 \leq i \leq n} w_i^S(X) \right] = 0.$$

Then W is consistent.

There are two main differences between Theorem 12 and Stone's theorem.

1. Conditions 1. and 3. of Stone's theorem are no longer necessary. This is because r -separated distributions are well-separated and thus have simpler conditions for consistency. In fact, a slight modification of the arguments of (Stone, 1977) shows that for r -separated distributions, condition 2. alone is sufficient for consistency.
2. Condition 2. is strengthened. Instead of requiring the weight of x_i 's outside of a given radius to go to 0 for $X \sim \mathcal{D}$, we require the same to *uniformly* hold over a ball centered at X .

Theorem 12 provides a general condition that allows us to verify the r -consistency of non-parametric methods. We now show below that two common non-parametric algorithms – k_n -nearest neighbors and kernel classifiers with rapidly decaying kernel functions – satisfy the conditions of Theorem 12.

Corollary 14. *Let $\mathcal{D}$ be any r -separated distribution. Let k_n be any sequence such that $\lim_{n \rightarrow \infty} \frac{k_n}{n} = 0$, and let M be the k_n -nearest neighbors classifier on a sample $S \sim \mathcal{D}^n$. Then M is r -consistent with respect to $\mathcal{D}$.*

Remarks:

1. Because the data distribution is r -separated, $k_n = 1$ will be r -consistent. Also observe that for r -separated distributions, $k_n = 1$ will converge towards the Bayes Optimal classifier.
2. In general, M converges towards the Bayes Optimal classifier provided that $k_n \rightarrow \infty$ in addition to $k_n/n \rightarrow 0$. This condition is not necessary for r -consistency – because the distribution is r -separated.

We next show that kernel classifiers are also r -consistent on r -separated data distributions, provided the kernel function decreases rapidly enough.

Corollary 15. *Let W be a kernel classifier over $\mathcal{X} \times \{\pm 1\}$ constructed from K and h_n . Suppose the following properties hold for K and h_n .*

1. *For any $c > 1$, $\lim_{x \rightarrow \infty} \frac{K(cx)}{K(x)} = 0$.*

2. $\lim_{n \rightarrow \infty} h_n = 0$.

If $\mathcal{D}$ is an r -separated distribution over $\mathcal{X} \times \{\pm 1\}$, then W is r -consistent with respect to $\mathcal{D}$.

Observe that Condition 1. is satisfied for any $K(x)$ that decreases more rapidly than an inverse polynomial – and is hence satisfied by most popular kernels like the Gaussian kernel. Is the condition on K in Corollary 15 necessary? The following example illustrates that a kernel classifier with any arbitrary K is not necessarily r -consistent. This indicates that some sort of condition needs to be imposed on K to ensure r -consistency; finding a tight necessary condition however is left for future work.

Example 3 Let $\mathcal{X} = [-1, 1]$ and let $\mathcal{D}$ be a distribution with $p_{\mathcal{D}}(-1, -1) = 0.1$ and $p_{\mathcal{D}}(1, 1) = 0.9$. Clearly, $\mathcal{D}$ is 0.3-separated. Let $K(x) = e^{-\min(|x|, 0.2)^2}$. Let h_n be any sequence with $\lim_{n \rightarrow \infty} h_n = 0$ and $\lim_{n \rightarrow \infty} nh_n = \infty$. Let W be the weight classifier with input $S = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$ such that

$$w_i^S(x) = \frac{K\left(\frac{|x-x_i|}{h_n}\right)}{\sum_{j=1}^n K\left(\frac{|x-x_j|}{h_n}\right)}.$$

W can be shown to satisfy all the conditions of Theorem 13 (the proof is analogous to the case for a Gaussian Classifier), and is therefore consistent. However, W does not learn a robust classifier on $\mathcal{D}$ for $r = 0.3$.

Consider $x = -0.7$. For any $\{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\} \sim \mathcal{D}^n$, all x_i will either be -1 or 1 . Therefore, since $K(|x - (-1)|) = K(|x - 1|)$, it follows that $w_i^S(x) = \frac{1}{n}$ for all $1 \leq i \leq n$. Since $x_i = 1$ with probability 0.9, it follows that with high probability x will be classified as 1 which means that f , the output of W , is not robust at $x = -1$. Thus f has astuteness at most 0.9 which means that W is not r -consistent for $r = 0.3$.

4. General Distributions

We next consider more general data distributions, where data from different classes may be close together in space, and may even overlap. Observe that unlike the r -separated case, here there may be no classifier with astuteness one. Thus, a natural question is: what does the optimally astute classifier look like, and how can we build non-parametric classifiers to this limit?

4.1. The r -Optimal Classifier and Adversarial Pruning

(Yang et al., 2019) propose a large-sample limit – called the r -optimal – and show that it is analogous to the Bayes

Optimal classifier for robustness. More specifically, given a data distribution D , to find the r -optimal classifier, we solve the following optimization problem.

$$\begin{aligned} \max_{S_{+1}, S_{-1}} & \int_{x \in S_{+1}} p(y = +1|x) d\mu_{\mathcal{D}}(x) + \\ & \int_{x \in S_{-1}} p(y = -1|x) d\mu_{\mathcal{D}}(x) \quad (1) \\ \text{subject to} & d(S_{+1}, S_{-1}) > 2r \end{aligned}$$

Then, the r -optimal classifier is defined as follows.

Definition 16. (Yang et al., 2019) Fix $r, \mathcal{D}$. Let S_{+1}^* and S_{-1}^* be any optimizers of (1). Then the r -optimal classifier, g_r^* is any classifier such that $g_r^*(x) = j$ whenever $d(S_j^*, x) \leq r$.

Yang et. al. (Yang et al., 2019) show that the r -optimal classifier achieves the optimal astuteness – out of all classifiers on the data distribution $\mathcal{D}$; hence, it is a robustness analogue to the Bayes Optimal Classifier. Therefore, for general distributions, the goal in robust classification is to find non-parametric algorithms that output classifiers that converge towards g_r^* .

To find robust classifiers, (Yang et al., 2019) propose Adversarial Pruning – a defense method that preprocesses the training data by making it better separated. More specifically, Adversarial Pruning takes as input a training dataset S and a radius r , and finds the largest subset of the training set where differently labeled points are at least distance $2r$ apart.

Definition 17. A set $S_r \subset \mathcal{X} \times \{\pm 1\}$ is said to be *r -separated* if for all $(x_1, y_1), (x_2, y_2) \in S_r$, if $y_1 \neq y_2$, then $d(x_1, x_2) > 2r$. To *adversarially prune* a set S is to return its largest r -separated subset. We let $\text{AdvPrun}(S, r)$ denote the result of adversarial pruning S .

Once an r -separated subset S_r of the training set is found, a standard non-parametric method is trained on S_r . While (Yang et al., 2019) show good empirical performance of such algorithms, no formal guarantees are provided. We next formally characterize when adversarial pruning followed by a non-parametric method results in a classifier that is provably r -consistent.

Specifically, we consider analysing the general algorithm provided in Algorithm 1.

4.2. Convergence Guarantees

We begin with some notation. For any weight function W and radius $r > 0$, we let $\text{RobustNonPar}(W, r)$ represent the weight function that outputs weights for $S \sim \mathcal{D}^n$

Algorithm 1 RobustNonPar

Input: $S \sim \mathcal{D}^n$, weight function W , robustness radius r
 $S_r \leftarrow \text{AdvPrun}(S, r)$

Output: W_{S_r}

according to $\text{RobustNonPar}(S, W, r)$. In particular, this can be used to convert any weight function algorithm into a new weight function which takes robustness into account. A natural question is, for which weight functions W is $\text{RobustNonPar}(W, r)$ r -consistent? Our next theorem provides sufficient conditions for this.

Theorem 18. *Let W be a weight function over $\mathcal{X} \times \{\pm 1\}$, and let $\mathcal{D}$ be a distribution over $\mathcal{X} \times \{\pm 1\}$. Fix $r > 0$. Let $S_r = \text{AdvPrun}(S, r)$. For convenience, relabel x_i, y_i so that $S_r = \{(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m)\}$. Suppose that for any $0 < a < b$,*

$$\lim_{n \rightarrow \infty} \mathbb{E}_{S \sim \mathcal{D}^n} \left[\frac{1}{m} \sum_{i=1}^m \sup_{x \in B(x_i, a)} \sum_{j=1}^m w_j^{S_r}(x) I_{\|x_j - x\| > b} \right] = 0.$$

Then $\text{RobustNonPar}(W, r)$ is r -consistent with respect to $\mathcal{D}$.

Remark: There are two important differences between the conditions in Theorem 18 and Theorem 12.

1. We replace S with S_r .
2. The expectation over $X \sim \mathcal{D}_{\mathcal{X}}$ is replaced with an average over $\{x_1, x_2, \dots, x_m\}$. The intuition here is that we are replacing $\mathcal{D}$ with a uniform distribution over S_r . While $\mathcal{D}$ may not be r -separated, the uniform distribution over S_r is, and represents the region of points where our classifier is astute.

A natural question is what satisfies the conditions in Theorem 18. We next show that k_n -nearest neighbors and kernel classifiers with rapidly decaying kernel functions continue to satisfy the conditions in Theorem 18; this means that these classifiers, when combined with Adversarial Pruning, will converge to r -optimal classifiers in the large sample limit.

Corollary 19. *Let k_n be a sequence with $\lim_{n \rightarrow \infty} \frac{k_n}{n} = 0$, and let M denote the k_n -nearest neighbor algorithm. Then for any $r > 0$, $\text{RobustNonPar}(M, r)$ is r -consistent.*

Remark: Corollary 19 gives a formal guarantee in the large sample limit for the modified nearest-neighbor algorithm proposed by (Yang et al., 2019).

Corollary 20. *Let W be a kernel classifier over $\mathcal{X} \times \{\pm 1\}$ constructed from K and h_n . Suppose the following properties hold for K and h_n .*

1. For any $c > 1$, $\lim_{x \rightarrow \infty} \frac{K(cx)}{K(x)} = 0$.
2. $\lim_{n \rightarrow \infty} h_n = 0$.

Then for any $r > 0$, $\text{RobustNonPar}(W, r)$ is r -consistent.

Observe again that Condition 1. is satisfied by any K that decreases more rapidly than an inverse polynomial kernel; it is thus satisfied by most popular kernels, such as the Gaussian kernel.

5. Validation

Our theoretical results are, by nature, large sample; we next validate how well they apply to the finite sample case by trying them out on a simple example. In particular, we ask the following question:

How does the robustness of non-parametric classifiers change with increasing sample size?

This question is considered in the context of two simple non-parametric classifiers – one nearest neighbor (which is guaranteed to be r -consistent) and histograms (which is not). To be able to measure performance with increasing data size, we look at a simple synthetic dataset – the Half Moons.

5.1. Experimental Setup

Classifiers and Dataset. We consider two different classification algorithms – one nearest neighbor (NN) and a Histogram Classifier (HC). We use the Halfmoon dataset with two settings of the gaussian noise parameter σ , $\sigma = 0$ (Noiseless) and $\sigma = 0.08$ (Noisy). For the Noiseless setting, observe that the data is already 0.1-separated; for the Noisy setting, we use Adversarial Pruning (Algorithm 1) with parameter $r = 0.1$ for both classification methods.

Performance Measure. We evaluate robustness with respect to the ℓ_∞ metric, that is commonly used in the adversarial examples literature. Specifically, for each classifier, we calculate the *empirical astuteness*, which is the fraction of test examples on which it is astute.

Observe that computing the empirical astuteness of a classifier around an input x amounts to finding the adversarial example that is *closest to x* according to the ℓ_∞ norm. For the 1-nearest neighbor, we do this using the optimal attack algorithm proposed by Yang et. al. (Yang et al., 2019). For the histogram classifier, we use the optimal attack framework proposed by (Yang et al., 2019), and show that the structure of the classifier can be exploited to solve the convex program efficiently. Details are in Appendix C.

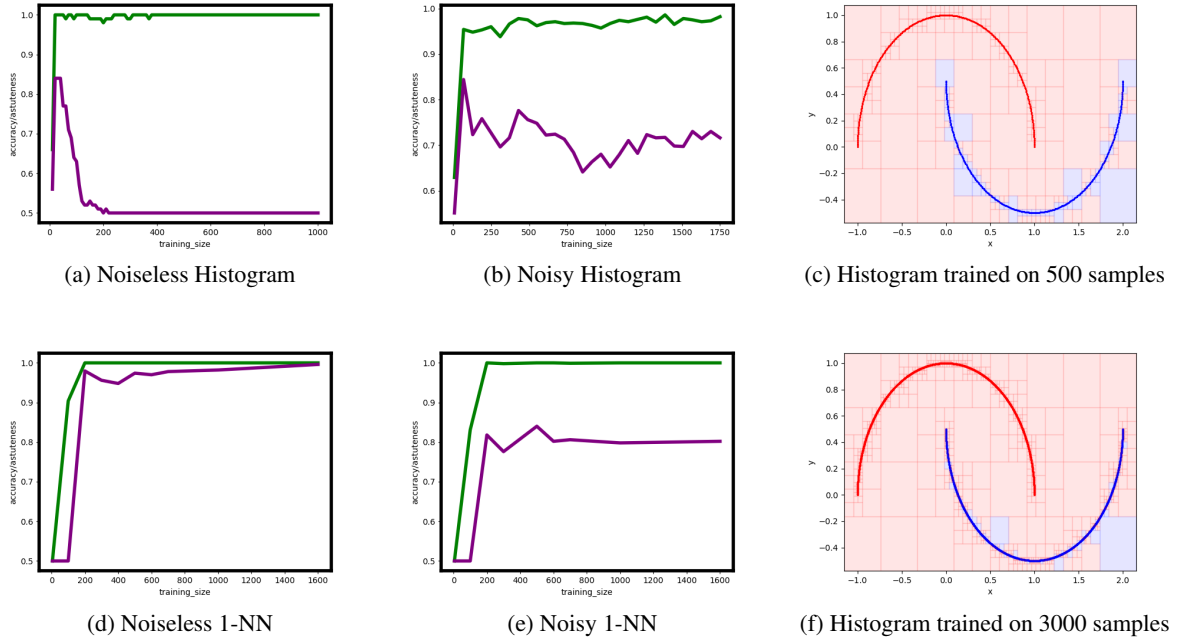


Figure 2. Empirical accuracy/astuteness of different classifiers as a function of training sample size. Left : Noiseless Setting. Right: Noisy Setting. Top Row: Histogram Classifier, Bottom Row: 1-Nearest Neighbor

We use an attack radius of $r = 0.1$ for the Noiseless setting, and $r = 0.09$ for the Noisy setting. For all classification algorithms, we plot the empirical astuteness as a function of the training set size. As a baseline, we also plot their standard accuracy on the test set.

5.2. Results

The results are presented in Figure 2; the left two panels are for the Noiseless setting while the two center ones are for the Noisy setting.

The results show that as predicted by our theory, for the Noiseless setting, the empirical astuteness of nearest neighbors converges to 1 as the training set grows. For Histogram Classifiers, the astuteness converges to 0.5 – indicating that the classifier may grow less and less astute with higher sample size even for well-separated data. This is plausibly because the cell size induced by the histogram grows smaller with growing training data; thus, the classifier that outputs the default label -1 in empty cells is incorrect on adversarial examples that are close to a point with $+1$ label, but belongs to a different, empty cell. The rightmost panels in Figure 2 provide a visual illustration of this process.

For the Noisy setting, the empirical astuteness of adversarial pruning followed by nearest neighbors converges to 0.8. For histograms with adversarial pruning, the astuteness converges to 0.7, which is higher than the noiseless case but still clearly sub-optimal.

5.3. Discussion

Our results show that even though our theory is asymptotic, our predictions continue to be relevant in finite sample regimes. In particular, on well-separated data, nearest neighbors that we theoretically predict to be intrinsically robust is robust; histogram classifiers, which do not satisfy the conditions in Theorem 12 are not. Our predictions continue to hold for data that is not well-separated. Nearest neighbors coupled with Adversarial Pruning continues to be robust with growing sample size, while histograms continue to be non-robust. Thus our theory is confirmed by practice.

6. Conclusion

In conclusion, we rigorously analyze when non-parametric methods provide classifiers that are robust in the large sample limit. We provide a general condition that characterizes when non-parametric methods are robust on well-separated data, and show that Adversarial Pruning of (Yang et al., 2019) works on data that is not well-separated.

Our results serve to provide a set of guidelines that can be used for designing non-parametric methods that are robust and accurate on well-separated data; additionally, we demonstrate that when data is not well-separated, preprocessing by adversarial pruning (Yang et al., 2019) does lead to optimally astute solutions in the large sample limit.

Acknowledgements

We thank NSF under CNS 1804829 for research support.

References

- Amsaleg, L., Bailey, J., Barbe, D., Erfani, S. M., Houle, M. E., Nguyen, V., and Radovanovic, M. The vulnerability of learning to adversarial perturbation increases with intrinsic dimensionality. In *2017 IEEE Workshop on Information Forensics and Security, WIFS 2017, Rennes, France, December 4-7, 2017*, pp. 1–6, 2017.
- Andriushchenko, M. and Hein, M. Provably robust boosted decision stumps and trees against adversarial attacks. In Wallach, H., Larochelle, H., Beygelzimer, A., Fox, E., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems 32*, pp. 12997–13008. Curran Associates, Inc., 2019.
- Avdiukhin, D., Mitrovic, S., Yaroslavlsev, G., and Zhou, S. Adversarially robust submodular maximization under knapsack constraints. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining, KDD 2019, Anchorage, AK, USA, August 4-8, 2019*, pp. 148–156, 2019. doi: 10.1145/3292500.3330911.
- Carlini, N. and Wagner, D. A. Towards evaluating the robustness of neural networks. In *2017 IEEE Symposium on Security and Privacy, SP 2017, San Jose, CA, USA, May 22-26, 2017*, pp. 39–57, 2017.
- Chen, H., Zhang, H., Boning, D. S., and Hsieh, C.-J. Robust decision trees against adversarial examples. *CoRR*, abs/1902.10660, 2019.
- Devroye, L., Györfi, L., and Lugosi, G. *A Probabilistic Theory of Pattern Recognition*, volume 31 of *Stochastic Modelling and Applied Probability*. Springer, 1996.
- Gates, G. W. The reduced nearest neighbor rule (corresp.). *IEEE Trans. Information Theory*, 18(3):431–433, 1972.
- Goodfellow, I. J., Shlens, J., and Szegedy, C. Explaining and harnessing adversarial examples, March 20 2014. URL <http://arxiv.org/abs/1412.6572>.
- Gottlieb, L., Kontorovich, A., and Nisnevitch, P. Near-optimal sample compression for nearest neighbors. In *Advances in Neural Information Processing Systems 27: Annual Conference on Neural Information Processing Systems 2014, December 8-13 2014, Montreal, Quebec, Canada*, pp. 370–378, 2014.
- Hanneke, S., Kontorovich, A., Sabato, S., and Weiss, R. Universal bayes consistency in metric spaces. *CoRR*, abs/1906.09855, 2019.
- Hart, P. E. The condensed nearest neighbor rule (corresp.). *IEEE Trans. Information Theory*, 14(3):515–516, 1968.
- Hein, M. and Andriushchenko, M. Formal guarantees on the robustness of a classifier against adversarial manipulation. In Guyon, I., Luxburg, U. V., Bengio, S., Wallach, H., Fergus, R., Vishwanathan, S., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems 30*, pp. 2266–2276. Curran Associates, Inc., 2017.
- Kantchelian, A., Tygar, J. D., and Joseph, A. D. Evasion and hardening of tree ensemble classifiers. *CoRR*, abs/1509.07892, 2015. URL <http://arxiv.org/abs/1509.07892>.
- Katz, G., Barrett, C. W., Dill, D. L., Julian, K., and Kochenderfer, M. J. Towards proving the adversarial robustness of deep neural networks. In *Proceedings First Workshop on Formal Verification of Autonomous Vehicles, FVAV@iFM 2017, Turin, Italy, 19th September 2017*, pp. 19–26, 2017.
- Kontorovich, A. and Weiss, R. A bayes consistent 1-nn classifier. In *Proceedings of the Eighteenth International Conference on Artificial Intelligence and Statistics, AISTATS 2015, San Diego, California, USA, May 9-12, 2015*, 2015.
- Kontorovich, A., Sabato, S., and Weiss, R. Nearest-neighbor sample compression: Efficiency, consistency, infinite dimensions. In *Advances in Neural Information Processing Systems 30: Annual Conference on Neural Information Processing Systems 2017, 4-9 December 2017, Long Beach, CA, USA*, pp. 1573–1583, 2017.
- Liu, Y., Chen, X., Liu, C., and Song, D. Delving into transferable adversarial examples and black-box attacks. In *5th International Conference on Learning Representations, ICLR 2017, Toulon, France, April 24-26, 2017, Conference Track Proceedings*, 2017.
- Lowd, D. and Meek, C. Adversarial learning. In Grossman, R., Bayardo, R. J., and Bennett, K. P. (eds.), *Proceedings of the Eleventh ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, Chicago, Illinois, USA, August 21-24, 2005*, pp. 641–647. ACM, 2005. ISBN 1-59593-135-X.
- Madry, A., Makelov, A., Schmidt, L., Tsipras, D., and Vladu, A. Towards deep learning models resistant to adversarial attacks. In *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*, 2018.
- Montasser, O., Hanneke, S., and Srebro, N. VC classes are adversarially robustly learnable, but only improperly. In

- Conference on Learning Theory, COLT 2019, 25-28 June 2019, Phoenix, AZ, USA*, pp. 2512–2530, 2019.
- Papernot, N., McDaniel, P. D., Jha, S., Fredrikson, M., Celik, Z. B., and Swami, A. The limitations of deep learning in adversarial settings. In *IEEE European Symposium on Security and Privacy, EuroS&P 2016, Saarbrücken, Germany, March 21-24, 2016*, pp. 372–387, 2016a.
- Papernot, N., McDaniel, P. D., Wu, X., Jha, S., and Swami, A. Distillation as a defense to adversarial perturbations against deep neural networks. In *IEEE Symposium on Security and Privacy, SP 2016, San Jose, CA, USA, May 22-26, 2016*, pp. 582–597, 2016b.
- Papernot, N., McDaniel, P. D., Goodfellow, I. J., Jha, S., Celik, Z. B., and Swami, A. Practical black-box attacks against deep learning systems using adversarial examples. *ASIACCS*, 2017.
- Raghunathan, A., Steinhardt, J., and Liang, P. Certified defenses against adversarial examples. In *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*, 2018.
- Salman, H., Yang, G., Li, J., Zhang, P., Zhang, H., Razenshteyn, I. P., and Bubeck, S. Provably robust deep learning via adversarially trained smoothed classifiers. *CoRR*, abs/1906.04584, 2019. URL <http://arxiv.org/abs/1906.04584>.
- Schmidt, L., Santurkar, S., Tsipras, D., Talwar, K., and Madry, A. Adversarially robust generalization requires more data. In *Advances in Neural Information Processing Systems 31*, pp. 5014–5026. Curran Associates, Inc., 2018.
- Sinha, A., Namkoong, H., and Duchi, J. C. Certifying some distributional robustness with principled adversarial training. In *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*, 2018.
- Sitawarin, C. and Wagner, D. A. On the robustness of deep k-nearest neighbors. In *2019 IEEE Security and Privacy Workshops, SP Workshops 2019, San Francisco, CA, USA, May 19-23, 2019*, pp. 1–7, 2019.
- Steinwart, I. Consistency of support vector machines and other regularized kernel classifiers. *IEEE Trans. Information Theory*, 51(1):128–142, 2005.
- Stone, C. Consistent nonparametric regression. *The Annals of Statistics*, 5(4):595–645, 1977.
- Szegedy, C., Zaremba, W., Sutskever, I., Bruna, J., Erhan, D., Goodfellow, I. J., and Fergus, R. Intriguing properties of neural networks. In *2nd International Conference on Learning Representations, ICLR 2014, Banff, AB, Canada, April 14-16, 2014, Conference Track Proceedings*, 2014.
- Wang, Y., Jha, S., and Chaudhuri, K. Analyzing the robustness of nearest neighbors to adversarial examples. In *Proceedings of the 35th International Conference on Machine Learning, ICML 2018, Stockholmsmässan, Stockholm, Sweden, July 10-15, 2018*, pp. 5120–5129, 2018.
- Yang, Y., Rashtchian, C., Wang, Y., and Chaudhuri, K. Adversarial examples for non-parametric methods: Attacks, defenses and large sample limits. *CoRR*, abs/1906.03310, 2019. URL <http://arxiv.org/abs/1906.03310>.

A. Proofs for r -separated distributions

For any distribution $\mathcal{D}$ over $\mathcal{X} \times Y$, it will be convenient to use the following notation: for any measurable $S \subset \mathcal{X}$, let $\mathbb{P}_{\mathcal{D}}[S] = \mathbb{P}_{(x,y) \sim \mathcal{D}}[x \in S]$. The following definition will be central to our proofs.

Definition 21. Let $\mathcal{D}$ be a distribution over $\mathcal{X} \times Y$. An $(\epsilon, \gamma, \alpha)$ -**decomposition** of $\mathcal{D}$ is a finite set of closed balls $B_1, B_2, \dots, B_s \subset \mathcal{X}$ each with radius γ such that

$$\mathbb{P}_{\mathcal{D}}[\cup_1^s B_i] > 1 - \epsilon,$$

and such that $\mathbb{P}_{\mathcal{D}}[B_i] \geq \alpha > 0$ for $1 \leq i \leq s$.

Lemma 22. Let $\mathcal{X}$ be a totally bounded metric space. For any distribution $\mathcal{D}$, and $\epsilon, \gamma > 0$, there exists $\alpha > 0$ such that $\mathcal{D}$ admits a $(\epsilon, \gamma, \alpha)$ -decomposition.

Proof. Fix any $x \in \mathcal{X}$ and $\epsilon, \gamma > 0$. Then the sequence of balls $\{S_i = B(x, i)\}$ has union equal to $\mathcal{X}$. Therefore, there exists j such that $\mathbb{P}_{\mathcal{D}}(S_j) > 1 - \epsilon$. Since S_j is totally bounded and complete, it is compact. Let $B^o(x, a)$ denote the open ball centered at x with radius a . Therefore, taking an open cover of S_j , $\{B^o(x, \gamma) : x \in S_j\}$, we can take a finite subcover $\{B_1^o, B_2^o, \dots, B_t^o\}$ that cover S_j . Discarding balls such that $\mathbb{P}_{\mathcal{D}}(B_i^o) = 0$ and taking the closure of each ball gives the desired result, with $\alpha = \min_i \mathbb{P}_{\mathcal{D}}(B_i)$. $\square$

To prove Theorem 12, we use the following lemma.

Lemma 23. Let $\mathcal{D}$ be a distribution over $\mathcal{X} \times \{\pm 1\}$, and let $B_1, B_2, \dots, B_s$ be a $(\epsilon, \gamma, \alpha)$ -decomposition of $\mathcal{D}$, and let $r > 3\gamma$. If W is a weight function satisfying the conditions of Theorem 12, then for any $\delta > 0$ there exists N such that for $n \geq N$, with probability $1 - \delta$ over $S \sim \mathcal{D}^n$, and $w_1, w_2, \dots, w_n$ learned by W from S ,

$$\sup_{\{x: d(x, \cup_1^s B_i) \leq r - 3\gamma\}} \sum_1^n w_i(x) I_{d(x_i, x) > r} < \frac{1}{3}.$$

Proof. Fix $\delta > 0$, and let Y be the indicator variable defined as

$$Y = \begin{cases} 1 & \text{if } \sup_{\{x: d(x, \cup_1^s B_i) \leq r - 3\gamma\}} \sum_1^n w_i(x) I_{d(x_i, x) > r} \geq \frac{1}{3} \\ 0 & \text{if } \sup_{\{x: d(x, \cup_1^s B_i) \leq r - 3\gamma\}} \sum_1^n w_i(x) I_{d(x_i, x) > r} < \frac{1}{3} \end{cases}.$$

It suffices to show that there exists N such that for all $n \geq N$, $E_{S \sim \mathcal{D}}[Y] \leq \delta$.

Fix $S \sim \mathcal{D}^n$ and suppose that $Y = 1$. Then there exists x^*, B_i^* such that $d(x^*, B_i^*) \leq r - 3\gamma$ and such that

$$\sum_1^n w_i(x^*) I_{d(x_i, x^*) > r} \geq \frac{1}{3}.$$

By definition, B_i has radius γ , so by the triangle inequality, for any $x \in B_i^*$, $d(x, x^*) \leq 2\gamma + r - 3\gamma = r - \gamma$. This implies $x^* \in B(x, r - \gamma)$. Therefore, for any $x \in B_i^*$,

$$\sup_{x' \in B(x, r - \gamma)} \sum_1^n w_i(x') I_{d(x', x_i) > r} \geq \sum_1^n w_i(x^*) I_{d(x^*, x_i) > r} \geq \frac{1}{3}.$$

By the definition of an $(\epsilon, \gamma, \alpha)$ -decomposition, we have that $\mathbb{P}_{\mathcal{D}}(B_i^*) \geq \alpha$. As a consequence, we have that

$$\mathbb{E}_{X \sim \mathcal{D}_{\mathcal{X}}} \left[\sup_{x' \in B(X, r - \gamma)} \sum_1^n w_i(x') I_{\|x_i - x'\| > r} \right] \geq \mathbb{P}_{\mathcal{D}}[B_i^*] \frac{1}{3} \geq \frac{\alpha}{3}.$$

Since the previous inequality is guaranteed to hold if $Y = 1$, taking the expectation over S yields that

$$\mathbb{E}_{S \sim \mathcal{D}^n} \mathbb{E}_{X \sim \mathcal{D}_{\mathcal{X}}} \left[\sup_{x' \in B(X, r - \gamma)} \sum_1^n w_i(x') I_{\|x_i - x'\| > r} \right] \geq \frac{\alpha E[Y]}{3}.$$

By the conditions of Theorem 12, the left side of the equation must tend to 0 as $n \rightarrow \infty$. This implies that the same must hold for the right side. Therefore, $E[Y]$ tends to 0 as $n \rightarrow \infty$, and we can select N such that $E[Y] < \delta$ for $n \geq N$, which completes the proof. $\square$

Proof. (Theorem 12) Let W be a weight function that satisfies the condition of Theorem 12. Fix $\epsilon, \delta > 0$, and $\gamma < r/3$. Applying Lemma 22, let $B_1, B_2, \dots, B_s$ be an $(\epsilon, \gamma, \alpha)$ -decomposition of $\mathcal{D}$. Let T^+ and T^- be subsets of $\mathcal{X}$ corresponding to the definition of r -separation for $\mathcal{D}$.

For $S \sim \mathcal{D}^n$, let A denote the event that

$$\sup_{\{x: d(x, \cup_1^s B_i) \leq r-3\gamma\}} \sum_1^n w_i(x) I_{d(x_i, x) > r} < \frac{1}{3}.$$

Suppose A holds. Pick a B_i . Since T^+ and T^- have distance greater than $2r$, and $\text{diam}(B_i) \leq 2\gamma < r$, either $B_i \cap T^+ = \emptyset$ or $B_i \cap T^- = \emptyset$. Note that for n sufficiently large, both cannot be empty since $P_{\mathcal{D}}(B_i) \geq \alpha > 0$ and each x in the support of $\mathcal{D}$ is either in T^+ or T^- .

Without loss of generality, $B_i \cap T^- = \emptyset$. Then $B_i \cap T^+ \neq \emptyset$. B_i has diameter 2γ . Thus $d(B_i, T^-) > 2r - 2\gamma$. Let $x \in B(B_i, r - 3\gamma)$. Then if $(x_j, -) \in S$, by the triangle inequality, $d(x, x_j) > 2r - 2\gamma - (r - 3\gamma) = r + \gamma$.

Substituting this and using event A , we have that

$$\sum_1^n w_i^S(x) I_{(x_i, -) \in S} \leq \sum_1^n w_i^S(x) I_{d(x_i, x) > r} < \frac{1}{3}.$$

It follows that $W_S(x) = +1$. An analogous argument holds for $B_i \cap T^+ = \emptyset$. This implies that W_S is astute with radius $r - 3\gamma$ over all B_i .

$\cup B_i$ has measure at least $1 - \epsilon$. By Lemma 23, for any $\delta > 0$ event A holds with probability $1 - \delta$ for n sufficiently large. Therefore, for n sufficiently large, we see that $A_{r-3\gamma}(W_S, \mathcal{D}) \geq 1 - \epsilon$ with probability $1 - \delta$. Because ϵ, δ and γ were arbitrary, it follows that W is r -consistent, as desired. $\square$

Proof. (Corollary 14) For any $S = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\} \subset \mathcal{X} \times \{\pm 1\}$, let $w_i^S(x)$ be 1 if and only if x_i is one of the k_n nearest neighbors of x in the set $S_{\mathcal{X}} = \{x_1, x_2, \dots, x_n\}$. Let $\mathcal{D}$ be a distribution over $\mathcal{X} \times \{\pm 1\}$. By Theorem 12, it suffices to show that for any $0 < a < b$,

$$\lim_{n \rightarrow \infty} \mathbb{E}_{X \sim \mathcal{D}_{\mathcal{X}}} [\mathbb{E}_{S \sim \mathcal{D}^n} [\sup_{x' \in B(x, a)} \sum_1^n w_i^S(x') I_{d(x_i, x') > b}]] = 0.$$

Fix $0 < a < b$, and let $\epsilon > 0$.

Pick $\gamma > 0$ such that $a + 2\gamma < b$. This is possible for any $a < b$. Let $B_1, B_2, \dots, B_s$ be an $(\epsilon, \gamma, \alpha)$ -decomposition of $\mathcal{D}$. By applying a Chernoff bound followed by a union bound, for any $\delta > 0$ there exists n such that with probability $1 - \delta$ over $S \sim \mathcal{D}^n$, each B_i satisfies $|B_i \cap S_{\mathcal{X}}| \geq \frac{n\alpha}{2}$. Furthermore, if n is sufficiently large, then $\frac{n\alpha}{2} > k_n$ holds as well.

Consider any $x \in B_i$, and $x' \in B(x, a)$. B_i has radius γ and also satisfies $|B_i \cap S_{\mathcal{X}}| > k_n$. Therefore, there are at least k_n points within distance $a + 2\gamma$ of x . Because $a + 2\gamma < b$, it follows that none of the k_n nearest neighbors of x' can have distance more than b from x' . In particular,

$$\sum_1^n w_i^S(x') I_{d(x_i, x') > b} = 0.$$

Since B_i, x and x' were arbitrary, we have that for all $x \in \cup B_i$,

$$\sup_{x' \in B(x, a)} \sum_1^n w_i^S(x') I_{d(x_i, x') > b} \leq \begin{cases} 0 & |B_i \cap S_{\mathcal{X}}| \geq \frac{n\alpha}{2}, 1 \leq i \leq s \\ 1 & \text{otherwise} \end{cases}$$

Since $X \in \cup_1^s B_i$ with probability at least $1 - \epsilon$, and since $|B_i \cap S_{\mathcal{X}}| \geq \frac{n\alpha}{2}, 1 \leq i \leq s$ with probability at least $1 - \delta$, it follows that

$$\mathbb{E}_{X \sim \mathcal{D}} [\mathbb{E}_{S \sim \mathcal{D}^n} [\sup_{x' \in B(x, a)} \sum_1^n w_i^S(x') I_{d(x_i, x') > b}]] \leq (1 - \delta - \epsilon)0 + \delta + \epsilon = \delta + \epsilon,$$

which can be made arbitrarily small as ϵ and δ were arbitrary. Therefore, the limit as n approaches infinity is 0, as desired. $\square$

Proof. (Corollary 15) Let $\mathcal{D}$ be a distribution over $\mathcal{X} \times \{\pm 1\}$. By Theorem 12, it suffices to show that for any $0 < a < b$,

$$\lim_{n \rightarrow \infty} \mathbb{E}_{X \sim \mathcal{D}} [\mathbb{E}_{S \sim \mathcal{D}^n} [\sup_{x' \in B(x, a)} \sum_{i=1}^n w_i^S(x') I_{d(x_i, x') > b}]] = 0.$$

Fix $0 < a < b$, and let $\epsilon > 0$.

Pick $\gamma > 0$ be such that $a + 2\gamma < b$. Let $B_1, B_2, \dots, B_s$ be an $(\epsilon, \gamma, \alpha)$ -decomposition of $\mathcal{D}$. By applying a Chernoff bound, for any $\delta > 0$ there exists n such that with probability $1 - \delta$ over $S \sim \mathcal{D}^n$, each B_i satisfies $|B_i \cap S_{\mathcal{X}}| \geq \frac{n\alpha}{2}$.

Next, consider any $x_i, x_j \in S_{\mathcal{X}}$, and let x be a point such that $d(x_i, x) \leq a + 2\gamma$ and $d(x_j, x) > b$. Then we have that

$$\frac{w_j^S(x)}{w_i^S(x)} = \frac{K(\frac{d(x_j, x)}{h_n})}{K(\frac{d(x_i, x)}{h_n})}.$$

Because $b > a + 2\gamma$, $\frac{d(x_j, x)}{d(x_i, x)} > 1$. Therefore, since $\lim_{n \rightarrow \infty} h_n = 0$ and $\lim_{x \rightarrow \infty} \frac{K(cx)}{K(x)} = 0$ for $c > 1$, it follows that for any $\beta > 0$, there exists N such that for $n \geq N$,

$$\frac{w_j^S(x)}{w_i^S(x)} \leq \frac{\alpha\beta}{2}.$$

Fix any such β , and consider any x with $d(x, B_i) \leq a$. Then $d(x, x') \leq a + 2\gamma < b$ for any $x' \in B_i$. Recall that B_i contains at least $\frac{n\alpha}{2}$ points, and let $c = \min_{i, d(x_i, x) \leq a+2\gamma} w_i(x)$. Then it follows that

$$\begin{aligned} \sum_{i=1}^n w_i^S(x) I_{d(x_i, x) > b} &\stackrel{(a)}{=} \frac{\sum_{i=1}^n w_i^S(x) I_{d(x_i, x) > b}}{\sum_{i=1}^n w_i^S(x)} \\ &\stackrel{(b)}{\leq} \frac{\sum_{i=1}^n w_i^S(x) I_{d(x_i, x) > b}}{\sum_{i=1}^n w_i^S(x) I_{d(x_i, x) \leq a+2\gamma}} \\ &\stackrel{(c)}{\leq} \frac{nc \frac{\alpha\beta}{2}}{\frac{n\alpha}{2} c} \\ &= \beta \end{aligned}$$

(a) holds because the weights always sum to 1. (b) holds because we are reducing the denominator. (c) holds because there are at least $\frac{n\alpha}{2}$ points in B_i , with c being the minimum weight (stated above). The numerator is a result of the inequality shown above in which $w_j^S(x)/w_i^S(x) \leq \alpha\beta/2$ if $d(x_j, x) > b$ and $d(x_i, x) \leq a + 2\gamma$.

Using this, we get the following bound:

$$\sup_{x' \in B(x, a)} \sum_{i=1}^n w_i^S(x') I_{d(x_i, x') > b} \leq \begin{cases} \beta & x \in \cup_1^s B_i, |B_i \cap S_{\mathcal{X}}| \geq \frac{n\alpha}{2}, 1 \leq i \leq s \\ 1 & \text{otherwise} \end{cases}$$

Since $x \in \cup_1^s B_i$ with probability $1 - \epsilon$, and since $|B_i \cap S_{\mathcal{X}}| \geq \frac{n\alpha}{2}, 1 \leq i \leq s$ with probability $1 - \delta$, it follows that

$$\mathbb{E}_{X \sim \mathcal{D}} [\mathbb{E}_{S \sim \mathcal{D}^n} [\sup_{x' \in B(x, a)} \sum_{i=1}^n w_i^S(x') I_{d(x_i, x') > b}]] \leq (1 - \delta - \epsilon)\beta + \delta + \epsilon.$$

which can be made arbitrarily small as ϵ, β , and δ were arbitrary. Therefore, the limit as n approaches infinity is 0, as desired. $\square$

B. Proofs for general distributions

Lemma 24. Let $B_1, \dots, B_s$ be a $(\epsilon, \alpha, \gamma)$ decomposition of $\mathcal{D}$ over $\mathcal{X} \times \{\pm 1\}$. Let $U \subseteq [s]$. Then if $n \geq O(\frac{s^{2^{2s}} \log(1/\delta)}{\epsilon^2})$, then with probability at least $1 - \delta$, for all U we have:

$$\begin{aligned} |\mathbb{P}_{(x,y) \sim \mathcal{D}}[x \in \cup_{i \in U} B_i, y = +] - \mathbb{P}_{(x,y) \sim \mathcal{D}_S}[x \in \cup_{i \in U} B_i, y = +]| &\leq \epsilon, \\ |\mathbb{P}_{(x,y) \sim \mathcal{D}}[x \in \cup_{i \in U} B_i, y = -] - \mathbb{P}_{(x,y) \sim \mathcal{D}_S}[x \in \cup_{i \in U} B_i, y = -]| &\leq \epsilon. \end{aligned}$$

Proof. For any given $U \subseteq [s]$, by a Chernoff bound we have that

$$|\mathbb{P}_{(x,y) \sim \mathcal{D}}[x \in \cup_{i \in U} B_i, y = +] - \mathbb{P}_{(x,y) \sim \mathcal{D}_S}[x \in \cup_{i \in U} B_i, y = +]| > \epsilon$$

with probability at most $\frac{\delta}{2^{s+1}}$. Taking a union bound over all U , we see that with probability $1 - \frac{\delta}{2}$,

$$|\mathbb{P}_{(x,y) \sim \mathcal{D}}[x \in \cup_{i \in U} B_i, y = +] - \mathbb{P}_{(x,y) \sim \mathcal{D}_S}[x \in \cup_{i \in U} B_i, y = +]| \leq \epsilon$$

for all $U \subseteq [m]$. Applying the same to $y = -1$ and taking a union bound implies the result. $\square$

Lemma 25. Let M be a classification algorithm over $\mathcal{X} \times \{\pm 1\}$, $r > 0$ be a radius, and $\mathcal{D}$ be a distribution over $\mathcal{X} \times \{\pm 1\}$. Then for any ϵ, δ over $(0, 1)$, and for all γ over $(0, r/2)$, there exists N such that for $n \geq N$, with probability $1 - \delta$ over $S \sim \mathcal{D}^n$,

$$A_{r-\gamma}(M_S, \mathcal{D}) \geq A_r(M_S, \mathcal{D}_S) - \epsilon,$$

where $\mathcal{D}_S$ denotes the uniform distribution over S .

Proof. (**Lemma 25**) Fix $\epsilon, \delta > 0$ and $\gamma < r/2$. Applying Lemma 22, let $B_1, \dots, B_s$ be a $(\epsilon, \alpha, \gamma)$ decomposition of $\mathcal{D}$.

Let T be the subset of S such that M_S is astute at T with radius r . Define:

$$\begin{aligned} I_T^+ &= \{i | (x_j, +) \in T, x_j \in B_i\} \\ I_T^- &= \{i | (x_j, -) \in T, x_j \in B_i\}. \end{aligned}$$

Observe that $I_T^+ \cap I_T^- = \emptyset$. To see this, notice that B_i has radius $\gamma < r/2$. This implies that any $(x_j, +), (x_k, -) \in B_i$ would force M_S to not be astute at either of those points. Thus we can think of I_T^+ being the set of positively labeled balls, and I_T^- being the set of negatively labeled balls.

Let $B^+ = \cup_{i \in I_T^+} B_i$ and $B^- = \cup_{i \in I_T^-} B_i$. Our strategy will be to argue that M_S must be robust with radius $r - 2\gamma$ at $B^+ \cup B^-$, and then to observe that $\mathbb{P}_{\mathcal{D}}[(B^+, +)] + \mathbb{P}_{\mathcal{D}}[(B^-, -)]$ must be close to $A_r(M_S, \mathcal{D}_S)$.

Let $T_{\mathcal{X}} \subset \mathcal{X}$ denote the set of all x_i such that $(x_i, y_i) \in T$. By the definitions of $\mathcal{D}_S$ and T , we have that

$$\begin{aligned} A_r(M_S, \mathcal{D}_S) &= \frac{|T|}{n} \\ &= \frac{|T_{\mathcal{X}} \cap B^+|}{n} + \frac{|T_{\mathcal{X}} \cap B^-|}{n} + \frac{|T_{\mathcal{X}} \setminus (B^+ \cup B^-)|}{n}. \end{aligned}$$

If $x_i \in \cup_1^s B_j$ and $x_i \in T_{\mathcal{X}}$, then by definition, $x \in (B^+ \cup B^-)$. Therefore, $T_{\mathcal{X}} \setminus (B^+ \cup B^-)$ consists of $x_i \notin \cup_1^s B_j$. Using this, we see that

$$\begin{aligned} A_r(M_S, \mathcal{D}_S) &= \frac{|T_{\mathcal{X}} \cap B^+|}{n} + \frac{|T_{\mathcal{X}} \cap B^-|}{n} + \frac{|T_{\mathcal{X}} \setminus (B^+ \cup B^-)|}{n} \\ &\leq \mathbb{P}_{(x,y) \sim \mathcal{D}_S}[x \in B^+, y = +] + \mathbb{P}_{(x,y) \sim \mathcal{D}_S}[x \in B^-, y = -] + \mathbb{P}_{(x,y) \sim \mathcal{D}_S}[x \notin \cup_1^s B_j]. \end{aligned}$$

If n is sufficiently large, then by Lemma 24, each term on the right is within ϵ of its corresponding probability over $\mathcal{D}$. Thus we see that with probability $1 - \delta$,

$$A_r(M_S, \mathcal{D}_S) \leq \mathbb{P}_{(x,y) \sim \mathcal{D}}[x \in \cup_{i \in I_T^+} B_i, y = +] + \mathbb{P}_{(x,y) \sim \mathcal{D}}[x \in \cup_{i \in I_T^-} B_i, y = -] + 4\epsilon. \quad (2)$$

Observe that if M_S is robust with radius r at $x_j \in B_i$, then it is robust with radius $r - 2\gamma$ at all $x \in B_i$. Furthermore, for $x_j \in \cup_{i \in I_T^+} B_i$, M_S is astute at $(x_j, +1)$ with radius r . Therefore $M_S(x) = +1$ for all $x \in \cup_{i \in I_T^+} B_i$. Consequently,

$$\begin{aligned} A_{r-2\gamma}(M_S, \mathcal{D}) &\geq \mathbb{P}_{(x,y) \sim \mathcal{D}}[x \in \cup_{i \in I_T^+} B_i, y = +] + \mathbb{P}_{(x,y) \sim \mathcal{D}}[x \in \cup_{i \in I_T^-} B_i, y = -] \\ &\geq A_r(M_S, \mathcal{D}_S) - 4\epsilon \text{ (by equation 2).} \end{aligned}$$

Since this equation holds with probability $1 - \delta$, and since ϵ and γ were arbitrary, the result follows. $\square$

Proof. (Theorem 18) For convenience, we let W' represent the weight function described by $\text{RobustNonPar}(S, W, r)$. In particular, W'_S and W_{S_r} are the same classifier, where S_r denotes the largest r -separated subset of S .

Fix $\epsilon, \delta > 0$, and let $0 < \gamma < r$. For convenience, let

$$Z_i = \sup_{x \in B(x_i, r-\gamma)} \sum_{j=1}^m w_j^{S_r}(x) I_{||x_j - x|| > r}.$$

Because W fulfills the conditions of Theorem 18, there exists N such that for $n > N$, with probability $1 - \delta$ over $S \sim \mathcal{D}^n$, $\frac{1}{m} \sum_{i=1}^m Z_i < \epsilon$. Therefore, there exist at most $3m\epsilon$ values of i for which $Z_i > \frac{1}{3}$.

Since S_r is r -separated, it follows that

$$\sup_{x \in B(x_i, r-\gamma)} \sum_{j=1}^m w_j^{S_r}(x) I_{y_j \neq y_i} \leq Z_i.$$

Consequently, if $Z_i \leq \frac{1}{3}$, then $W_{S_r}(x) = y_i$ for all $x \in B(x_i, r - \gamma)$. Let $\mathcal{D}_S$ denote the uniform distribution over S . Then we have that

$$A_{r-\gamma}(W'_S, \mathcal{D}_S) = A_{r-\gamma}(W_{S_r}, \mathcal{D}_S) \geq \frac{|S_r|}{n} - 3\epsilon.$$

Observe that for n sufficiently large, with probability $1 - \delta$, $|A_r(g_r^*, \mathcal{D}) - A_r(g_r^*, \mathcal{D}_S)| \leq \epsilon$. The maximum possible astuteness over $\mathcal{D}_S$ is $\frac{|S_r|}{n}$ since no classifier can be astute at 2 oppositely labeled points with distance at most $2r$. Therefore, with probability $1 - 2\delta$,

$$A_{r-\gamma}(W'_S, \mathcal{D}_S) \geq A_r(g_r^*, \mathcal{D}) - 4\epsilon.$$

By Lemma 25, for n sufficiently large, with probability $1 - \delta$

$$A_{r-2\gamma}(W'_S, \mathcal{D}) \geq A_{r-\gamma}(W'_S, \mathcal{D}_S) - \epsilon.$$

Therefore, for n sufficiently large, with probability $1 - 3\delta$ over $S \sim \mathcal{D}$,

$$A_{r-2\gamma}(W'_S, \mathcal{D}) \geq A_r(g_r^*, \mathcal{D}) - 5\epsilon.$$

Since ϵ, δ , and γ were arbitrary, we are done. $\square$

The following two quick lemmas are used for the proofs of Corollaries 19 and 20.

Lemma 26. Let $B_1, B_2, \dots, B_s \subset \mathcal{X}$ denote s balls. Let $T \subset \mathcal{X}$ satisfy $|T \cap \cup_1^s B_i| = m$. Let

$$I_k \subseteq [s] = \{i : |B_i \cap T| \geq k\}.$$

Then $|\cup_{i \in I_k} B_i \cap T| \geq m - ks$.

Proof. For any $j \notin I_k$, $|B_j \cap T| < k$. Since there are at most s such j , it follows that $|\cup_{i \notin I_k} B_i \cap T| < ks$. Taking the complement implies the result. $\square$

Lemma 27. Let S be a finite subset of $\mathcal{X} \times \{\pm 1\}$. For any $r > 0$, let S_r denote the largest r -separated subset of S . Then $|S_r| \geq \frac{|S|}{2}$.

Proof. Let $S = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$. Define:

$$S_+ = \{(x_i, y_i) : y_i = +1\}$$

$$S_- = \{(x_i, y_i) : y_i = -1\}.$$

Observe that S_+ and S_- are both r -separated and have union S . Therefore one must have cardinality at least $\frac{|S|}{2}$, which implies the same about $|S_r|$. $\square$

Proof. (Corollary 19) For convenience, we let W' represent the weight function described by $RobustNonPar(S, W, r)$. In particular, W'_S and W_{S_r} are the same classifier, where S_r denotes the largest r -separated subset of S .

Relabel the points in S so that

$$S_r = \{(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m)\},$$

with $m \leq n$. We will also let $S_r^{\mathcal{X}} = \{x_1, x_2, \dots, x_m\}$.

By Theorem 18, it suffices to show that for any $0 < a < b$,

$$\lim_{n \rightarrow \infty} \mathbb{E}_{S \sim \mathcal{D}^n} \left[\frac{1}{m} \sum_{i=1}^m \sup_{x \in B(x_i, a)} \sum_{j=1}^m w_j^{S_r}(x) I_{d(x_i, x) > b} \right] = 0,$$

where w_j denote the weight functions corresponding to W . Fix $0 < a < b$, and let $\epsilon > 0$.

Pick $\gamma > 0$ be such that $a + 2\gamma < b$. Let $B_1, B_2, \dots, B_s$ be a $(\epsilon, \gamma, \alpha)$ decomposition of $\mathcal{D}$. By applying a Chernoff bound, for any $\delta > 0$ there exists n_0 such that for $n \geq n_0$, with probability $1 - \delta$ over $S \sim \mathcal{D}^n$,

$$|S_{\mathcal{X}} \cap \cup_1^s B_i| \geq (1 - 2\epsilon)n.$$

By Lemma 27, $\frac{m}{n} \geq \frac{1}{2}$. It follows that $|S_r^{\mathcal{X}} \cap \cup_1^s B_i| \geq m(1 - 4\epsilon)$.

Let

$$J = \{i : |B_i \cap S_r^{\mathcal{X}}| \geq m \frac{\epsilon}{s}\}.$$

By Lemma 26 it follows that $|S_r^{\mathcal{X}} \cap \cup_{i \in J} B_i| \geq m(1 - 4\epsilon) - m\epsilon = m(1 - 5\epsilon)$.

Next, observe that if n is sufficiently large, then

$$\frac{k_n}{m} \leq \frac{2k_n}{n} \leq \frac{\epsilon}{s}.$$

Therefore, $|B_i \cap S_r^{\mathcal{X}}| \geq k_n$ for $i \in J$.

Fix any B_j with $j \in J$, and consider x with $d(x, B_j) \leq a$. Then $d(x, x') \leq a + 2\gamma < b$ for any $x' \in B_j$. Therefore, since $|S_r^{\mathcal{X}} \cap B_i| \geq k_n$, all k_n -nearest neighbors of x have distance at most b to x . This implies that

$$\sum_1^m w_i^{S_r}(x) I_{d(x_i, x) > b} = 0.$$

For convenience, let

$$f(x_i) = \sup_{x \in B(x_i, a)} \sum_{j=1}^m w_j^{S_r}(x) I_{d(x, x_j) > b}.$$

For $x_i \in \cup_{j \in J} B_j$, any $x \in B(x_i, a)$ trivially satisfies $d(x, B_i) \leq a$. Therefore, $f(x_i) = 0$. Since $|S_r^{\mathcal{X}} \cap \cup_{j \in J} B_j| \geq m(1 - 5\epsilon)$, and $f(x_i) \leq 1$ for all $1 \leq i \leq m$, we have that

$$\begin{aligned} \frac{1}{m} \sum_1^m f(x_i) &= \frac{1}{m} \left(\sum_{x_i \in \cup_{i \in J} B_i} f(x_i) + \sum_{x_i \notin \cup_{i \in J} B_i} f(x_i) \right) \\ &\leq \frac{1}{m} (0 + 5\epsilon m(1)) \\ &= 5\epsilon. \end{aligned}$$

Since all of our equations hold with probability $1 - \delta$ over S for sufficiently large n , this last one does as well. Since this entire expression is always at most 1 (regardless of S), and since δ, ϵ were arbitrary, we have that

$$\lim_{n \rightarrow \infty} E_{S \sim \mathcal{D}^n} \left[\frac{1}{m} \sum_{i=1}^m f(x_i) \right] = 0,$$

which completes the proof. $\square$

Proof. (Corollary 20) For convenience, we let W' represent the weight function described by $\text{RobustNonPar}(S, W, r)$. In particular, W'_S and W_{S_r} are the same classifier, where S_r denotes the largest r -separated subset of S .

Relabel the points in S so that

$$S_r = \{(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m)\},$$

with $m \leq n$. We will also let $S_r^{\mathcal{X}} = \{x_1, x_2, \dots, x_m\}$.

By Theorem 18, it suffices to show that for any $0 < a < b$,

$$\lim_{n \rightarrow \infty} \mathbb{E}_{S \sim \mathcal{D}^n} \left[\frac{1}{m} \sum_{i=1}^m \sup_{x \in B(x_i, a)} \sum_{j=1}^m w_j^{S_r}(x) I_{d(x_i, x) > b} \right] = 0,$$

where w_j are the weight functions corresponding to W . Fix $0 < a < b$, and let $\epsilon > 0$.

Pick $\gamma > 0$ be such that $a + 2\gamma < b$. Let $B_1, B_2, \dots, B_s$ be a $(\epsilon, \gamma, \alpha)$ decomposition of $\mathcal{D}$. By applying a Chernoff bound, for any $\delta > 0$ there exists n_0 such that for $n \geq n_0$, with probability $1 - \delta$ over $S \sim \mathcal{D}^n$,

$$|S_{\mathcal{X}} \cap \cup_1^s B_i| \geq (1 - 2\epsilon)n.$$

By Lemma 27, $\frac{m}{n} \geq \frac{1}{2}$. It follows that $|S_r^{\mathcal{X}} \cap \cup_1^s B_i| \geq m(1 - 4\epsilon)$.

Let

$$J = \{i : |B_i \cap S_r^{\mathcal{X}}| \geq \frac{m\epsilon}{s}\}.$$

By Lemma 26, $|S_r^{\mathcal{X}} \cap \cup_{i \in J} B_i| \geq m(1 - 4\epsilon) - m\epsilon = m(1 - 5\epsilon)$.

Next, consider any $x_i, x_j \in S_r^{\mathcal{X}}$, and let x be a point such that $d(x_i, x) \leq a + 2\gamma$ and $d(x_j, x) > b$. Recall that W is constructed from kernel function K and window parameter h_n . We then have that

$$\frac{w_j^S(x)}{w_i^S(x)} = \frac{K\left(\frac{d(x_j, x)}{h_n}\right)}{K\left(\frac{d(x_i, x)}{h_n}\right)}. \quad (3)$$

Because $b > a + 2\gamma$, $\frac{d(x_j, x)}{d(x_i, x)} > 1$. Fix any $\beta > 0$. Because $\lim_{n \rightarrow \infty} h_n = 0$ and $\lim_{x \rightarrow \infty} \frac{K(cx)}{K(x)} = 0$ for $c > 1$, there exists N such that for $n \geq N$,

$$\frac{w_j^S(x)}{w_i^S(x)} \leq \frac{\beta\epsilon}{s}.$$

Fix B_j with $j \in J$, and consider x with $d(x, B_j) \leq a$. By the triangle inequality, $d(x, x') \leq a + 2\gamma$ for all $x' \in B_j$. Then we have the following,

$$\begin{aligned} \sum_1^m w_i^{S_r}(x) I_{d(x_i, x) > b} &\stackrel{(a)}{=} \frac{\sum_1^m w_i^{S_r}(x) I_{d(x_i, x) > b}}{\sum_1^m w_i^{S_r}(x)} \\ &\stackrel{(b)}{\leq} \frac{\sum_1^m w_i^{S_r}(x) I_{d(x_i, x) > b}}{\sum_{x_i \in B_j} w_i^{S_r}(x)} \\ &\stackrel{(c)}{\leq} \frac{m \sup_{x_i: d(x_i, x) > b} w_i^{S_r}(x)}{m\epsilon/s \inf_{x_i \in B_j} w_i^{S_r}(x)} \\ &\stackrel{(d)}{\leq} \frac{\beta\epsilon/s}{\epsilon/s} = \beta. \end{aligned} \quad (4)$$

Equation (a) holds because the total sum of weights is always 1, (b) because all weights are nonnegative, (c) because $|B_j \cap S_r^{\mathcal{X}}| \geq m\epsilon/s$, and (d) because of equation 3.

Let

$$Z_i = \sup_{x \in B(x_i, a)} \sum_{j=1}^m w_j^{S_r}(x) I_{d(x, x_j) > b}.$$

For $x_i \in \cup_1^t B_j$, any $x \in B(x_i, a)$ trivially satisfies $d(x, B_i) \leq a$. By equation 4, it follows that $Z_i \leq \beta$. Since $|\cup_{j \in J} B_j \cap S_r^{\mathcal{X}}| \geq m(1 - 5\epsilon)$ and $Z_i \leq 1$ for all $1 \leq i \leq m$, we have that

$$\begin{aligned} \frac{1}{m} \sum_1^m Z_i &= \frac{1}{m} \left(\sum_{x_i \in \cup_{j \in J} B_j} Z_i + \sum_{x_i \notin \cup_{j \in J} B_j} Z_i \right) \\ &\leq (1 - 5\epsilon)\beta + 5\epsilon. \end{aligned}$$

Since all of our equations hold with probability $1 - \delta$ over S for sufficiently large n , this last one does as well. Since this entire expression is always at most 1 (regardless of S), and since δ, ϵ, β were arbitrary, we have that

$$\lim_{n \rightarrow \infty} E_{S \sim \mathcal{D}^n} \left[\frac{1}{m} \sum_1^m Z_i \right] = 0,$$

which completes the proof. □

C. Experimental Details

C.1. Optimal attacks against histogram classifiers

Let H be a histogram classifier, and let (x, y) be any labeled example. Let $r > 0$ be some fixed robustness radius. Recall that an *adversarial example* against H at (x, y) is any x' such that $x' \in B(x, r)$ and $H(x') \neq y$. Note that if $H(x) \neq y$, then x itself is an adversarial example. Conversely, if H is astute at (x, y) with radius r , then no adversarial example exists.

For arbitrary classifiers, finding adversarial examples at a given point can be challenging. However, recent work (Yang et. al. 2019) has shown that for non-parametric classifiers, there are tractable methods for doing so. The key insight is that non-parametric classifiers can be construed as a partitioning of input space into convex cells, with each cell having a given label. For example, Figure 3 gives a visualization for these cells in a histogram classifier.

Because these cells are convex, finding an adversarial example for H at (x, y) (here x is a point in $\mathbb{R}^2$, and y is a label) amounts to finding the closest cell $c \in H$ to x such that $H(c) \neq y$. While Yang et. al. (Yang et. al. 2019) presents convex programming algorithms for doing this, the case of histograms in the ℓ_∞ metric is much simpler.

As stated in definition 10, a histogram partitions the input space into hypercubes by iteratively splitting each cube into 2^d cubes with half the length. Therefore, the cells of a histogram are all hypercubes of varying sizes. For cell c , let $s(c)$ denote the length of the cube that c corresponds to, and let $H(c)$ denote the label H assigns to c . The key observation is that c contains an adversarial example for (x, y) if and only if $d(c, x) \leq s(c)/2 + r$, and $H(c) \neq y$. This yields the following algorithm:

Algorithm 2 was further optimized by utilizing nearest-neighbor type algorithms to find the “closest” cells to x . This was done by grouping cells by their radii, and utilizing a separate nearest-neighbor data structure for all cells of a given radius.

Although this algorithm doesn’t have the same performance metrics as those presented in (Yang et. al. 2019), it was easily sufficient for computing the empirical astuteness for our experiments.

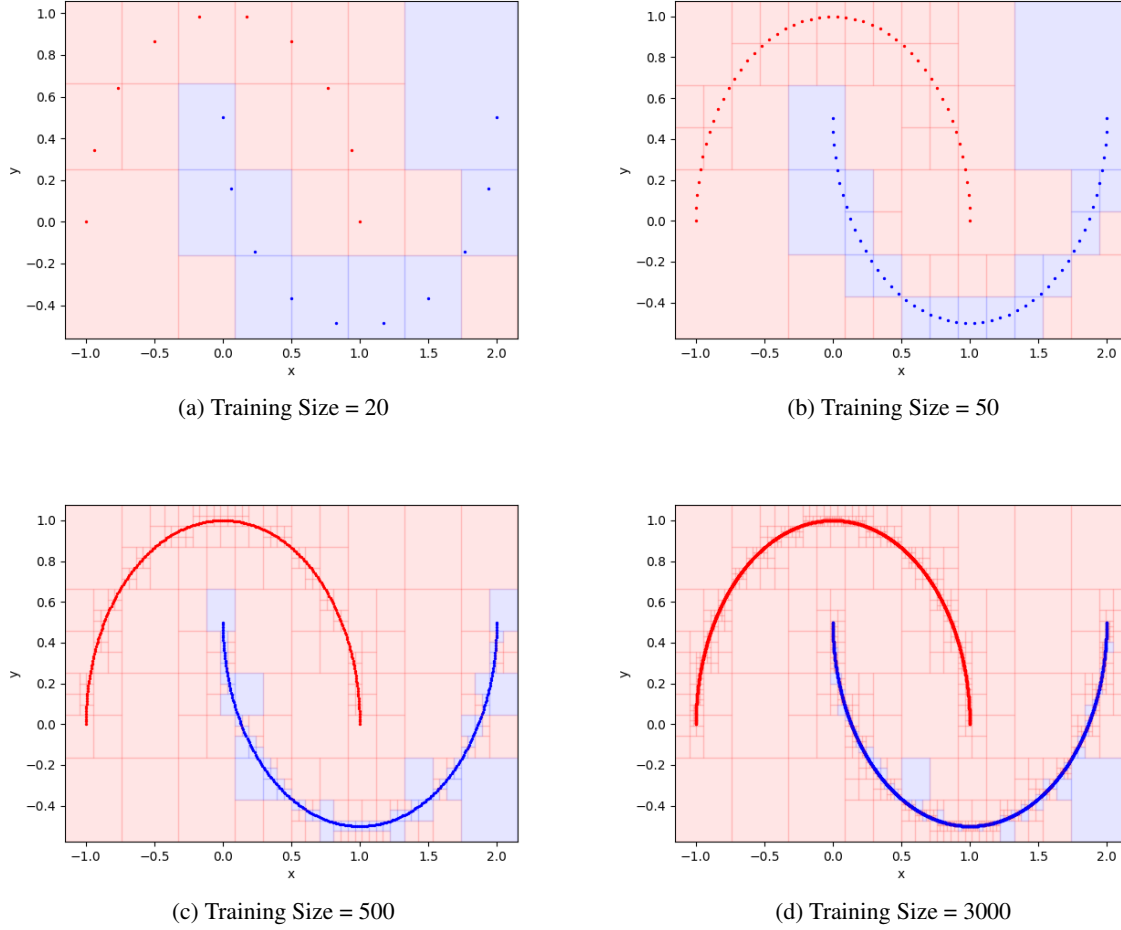


Figure 3. A visualization of histograms learned with training data sampled from noiseless halfmoon. As training size grows, the histogram classifier becomes increasingly susceptible to adversarial examples in the blue regions.

Algorithm 2 Optimal attack algorithm for Histogram Classifiers

Input: Histogram H , labeled point $(x, y) \in \mathbb{R}^2 \times \{\pm 1\}$, robustness radius r

for cell $c \in H$ **do**

if $d(c, x) \leq s(c)/2 + r$ and $H(c) \neq y$ **then**

Return c

end if

end for
